



Responsible artificial intelligence in human resources management: a review of the empirical literature

Antoine Bujold¹ · Isabelle Roberge-Maltais¹ · Xavier Parent-Rochelleau¹ · Jared Boasen¹ · Sylvain Sénécal¹ · Pierre-Majorique Léger¹

Received: 12 March 2023 / Accepted: 10 July 2023 / Published online: 19 July 2023
© The Author(s) 2023

Abstract

As it is the case for many business processes and activities disciplines, artificial intelligence (AI) is increasingly integrated in human resources management (HRM). While AI has great potential to augment the HRM activities in organizations, automating the management of humans is not without risks and limitations. The identification of these risks is fundamental to promote responsible use of AI in HRM. We thus conducted a review of the empirical academic literature across disciplines on the affordances and responsible principles of AI in HRM. This is the first review of responsible AI in HRM that focuses solely on studies containing observations, measurements, and tests about this phenomenon. The multi-domain and multidisciplinary approach and empirical focus provides a better understanding of the reality of the development, study, and deployment of AI in HRM and sheds light on how these are conducted responsibly. We conclude with a call for research based on what we identified as the most needed and promising avenues.

Keywords Human resources management · Artificial intelligence · Responsible AI · Human-centered AI

1 Introduction

Human resource management (HRM) activities comprise several routine and time-consuming tasks, while they are also subject to human perception, subjectivities, or biases. For these reasons, HRM is viewed as a fertile ground for the use of artificial intelligence [133, 143]. The use of artificial intelligence (AI) in HRM is being developed, tested, analyzed, and investigated empirically in various research domains [102, 125, 143]. Empirical investigations refer to studies based on data related to a phenomenon observed, measured and/or tested by the researchers [156]. Because there is no consensus on the definition of AI across and within domains due to the historical debate on what exactly “intelligence” is [44, 155] and AI is an umbrella term for different subset of technologies that mimic human intelligence (i.e., computer vision, natural language processing, machine learning, deep learning) [87, 144], this article will use a relatively broad, yet clear definition of the technology

that can be applied across the use of AI in HR. Specifically, in this paper AI is defined as “[...] the ability of a machine to learn from experience, adjust to new inputs and perform human-like tasks” [45, p. 63]. The rapid growth in the use of AI in HRM is reflected in the publication, in the last few years, of several literature and conceptual reviews on AI in HRM (e.g., [13, 20, 23, 34, 56, 67, 121, 123, 128, 145]).

Despite the important merits of these reviews, some limitations remain to our complete understanding of the affordances and risks of intelligent technologies in HRM, necessitating a thorough review of the literature from a different lens than the previous ones. Specifically, by focusing more on the literature of their respective domains, such as computer science or HRM, the previous reviews do not fully take into account the multi-domain nature of AI in HRM and the combination of both technical and social aspects of this phenomenon. Our study overcomes those limits by looking across domains at both (1) how AI is used in HRM (i.e., technical aspect) and (2) the responsible AI principles applied in our sample of studies (i.e., a social aspect).

Regarding the technical aspect, there is a certain lack of specificity on the technology studied (AI-enabled HRM). Specifically, because the reviews often fail to state explicitly and define what is the technology under examination,

✉ Antoine Bujold
antoine.bujold@hec.ca

¹ HEC Montreal, 3000 Chem. de la Côte-Sainte-Catherine, Montreal, QC H3T 2A7, Canada

recent reviews have described studies about various and not-necessarily AI-related technologies used in HRM (e.g., big data analysis, which is a massive amount of fine-grained and exhaustive data, but AI software is not ipso facto used to leverage this data [11, 74]). Our current review overcomes this limitation by including only studies that explicitly examine the use of AI, following the aforementioned definition, and thus clarifying to technical aspect of AI use in different HRM functions.

As for the social aspect, there is no current review with a focus on the responsible AI principles applied to HRM. Precisely, current reviews taking into account this social aspect mainly discuss or propose conceptual frameworks providing solutions on how AI should be studied, implemented or used, but none of them empirically observe the actual application of such frameworks. Our study contributes to knowledge by taking an inside look at how responsible principles are applied when developing, studying and deploying AI in HRM. Moreover, there is a necessity to look at the application of responsible research practices as many studies emphasize that responsibility is a key element when studying the use of AI in HRM (e.g., [6, 16, 61, 93, 147, 152, 153]). To our knowledge, this is the first systematic literature review looking precisely at what or which principles constitute responsible AI in HRM and how they are applied in empirical studies across domains.

However, as the notion of responsible use of technology is in constant evolution in the literature, there is no consensus around the definition and applications of responsible AI in the HRM domain. In this study we will adapt the broad definition of responsible AI from Barredo Arrieta et al. [19] which states that it is “[...] a series of AI principles to be necessarily met when deploying AI in real applications” [19, p. 83]. This adaptation will be done by including the responsible way of studying AI and defining responsible AI as a set of ethics principles to be necessarily followed when developing, studying, and deploying AI [133]. This definition will guide our review, but also provides researchers, organizations, and policy-makers of the necessary common understanding of what responsible AI refers to.

In sum, the aim of this article is to examine the scope of the existing empirical literature on responsible AI in HRM while attempting to overcome the limitations of previous work by conducting a systematic literature review including only empirical studies, all types of journals (not just in HRM), and no a priori conceptual framework. The objectives of this review are to: (1) identify empirical studies of current uses of AI in HRM, (2) review empirical knowledge of responsible AI principles in HRM and their application, and (3) evaluate the extent to which these research practices promote the combination of AI use with ethical, dignifying and quality work.

2 Methodology

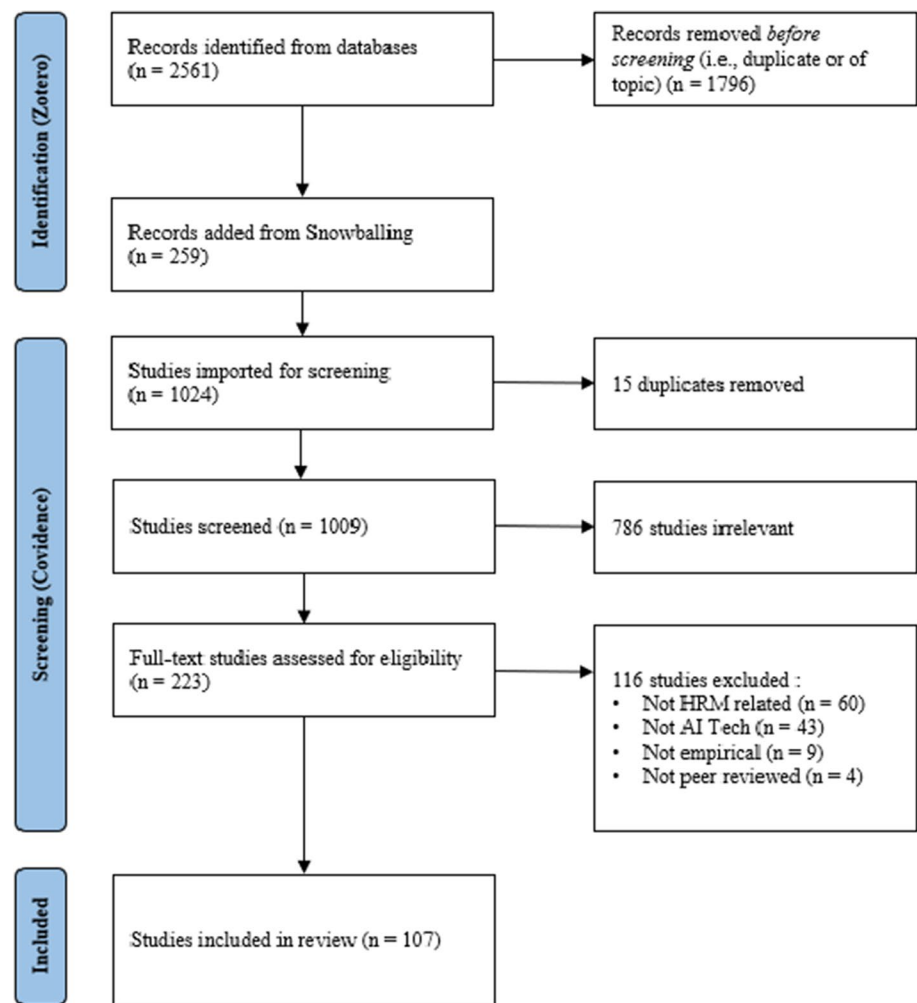
2.1 Retrieval

To guide our review, we followed the PRISMA 2020 statement, which allows for transparent reporting of our search strategy and our findings [118]. To be included, articles had to: (1) be an empirical study, (2) be peer-reviewed, (3) explicitly be related to a human resource management function, and (4) explicitly include an AI-driven technology based on the definition of AI presented in the introduction. To identify studies, we searched the following databases: Academic search complete, Business Source complete, PsycArticles, Web of science, and ABI/INFORM Collection. The broad scope and variety of these databases allowed us to assess multiple research domains in our literature review.

Appendix 1 presents the search query looking at the intersection of three areas. The first section includes domain related terms (e.g., human resource), the second includes responsible practices related terms (e.g., responsible or business ethics), and the third includes AI related terms (e.g., machine learning). The keywords included in our query were found using a two-step method commonly used in reviews [4, 122, 146]. The first step was to use the following search structure: Domain related terms “AND” responsible practice related terms. This search was conducted in each database. Fifty random studies per database were screened (abstract and title) to deduce any additional search terms that may have been missed. The second step is to use the following search structure: Domain related terms “AND” AI related terms. This was again searched in each database, with a maximum of 50 random studies reviewed [146].

At this point the number of records returned was 2561. The references were organized with the bibliography manager Zotero (Corporation for Digital Scholarship)¹ and the data was managed with Covidence (Covidence inc., Australia), an online platform for managing systematic reviews, and multiple spreadsheets. Duplicates were automatically detected by Zotero and deleted manually. Off-topic records (e.g., in the veterinary field) were also deleted manually for a total of 1796 removed records, leaving 765 records remaining. We then used the “snowball” approach to add more records that appeared to have relevant titles ($n=259$). The snowballing technique is used to enrich systematic reviews by using the references of articles in their existing samples to identify other articles potentially relevant to the reviews [159]. This technique was particularly important for our review because the literature on AI in HRM is rapidly evolving, and freshly published work or conference proceedings

¹ <https://www.zotero.org/>.

Fig. 1 PRISMA flow diagram

may have been slow to enter the databases we searched. The 1024 identified records were then transferred to Covidence, which automatically deleted the remaining duplicates that had bypassed the first process ($n = 15$).

Hence, 1009 records were identified for the title and abstract screening phase in Covidence. To ensure concordance before proceeding with record screening, an inter-coder reliability score was calculated using the percentage agreement (we agreed in advance that if it reached $> 75\%$, we would move on) [118, 146]. Specifically, in a pilot test, two researchers independently reviewed the title and abstract of a random sample of 50 records based on the four selection criteria and specified which criteria were not met if the study was excluded [146]. Their work was then compared. Only one round of pilot testing was required, with both researchers screening 42 of the 50 records in exactly the same way (i.e., a score of 84%).

The title and abstract of the 1009 records were then screened. Based on the selection criteria, 786 studies were deemed not relevant for the literature review. We then performed a full-text review of the remaining 223 studies and

excluded 116 which did not conform to the selection criteria (e.g., not HRM-related or AI-focused). In the end, a total of 107 studies were included in this systematic review. Figure 1 shows our PRISMA flow diagram.

2.2 Data extraction

First, after a thorough reading, two members of the research team listed each of the texts in detail in a summary spreadsheet, recording various characteristics related to the manuscript and the reported results. Theses syntheses were manually compared and found to be highly similar. The rare dissimilarities that occurred were resolved through discussion within the research team. Table 1 shows the data extraction categories used.

Regarding the meta-category about the use of AI in HRM, the researchers followed a recent conceptualization of algorithmic HRM by Meijerink and Bondarouk [101] as a guide to identify HRM functions from the data in the analyzed summary table. Meijerink and Bondarouk [101] describe the affordances of AI algorithms as talent

Table 1 Data extraction categories

Meta-categories	Categories
Study design	Data type (e.g., quantitative, qualitative) Type of study (e.g., experiment, field study)
Context	Aim of study Keywords Sector of activity Country Methodology Sample size Theoretical background
Human resources management	HRM functions HRM algorithm type
Responsible AI	Inclusion of responsible practices (yes/no) Type of responsible practices (if applicable)
Results	Conclusion and results General comments of the team member

acquisition, performance evaluation, talent management, workforce planning, and compensation and benefits. Moreover, to bring greater clarity and detail in our analysis of the technical aspect of AI in HRMR, we further granularized this meta-category according to whether the associated AI algorithms were descriptive, predictive, and/or prescriptive, as per the work of Leicht-Deobald et al. [90]. Thus, those meta-categories and granularized sub-categories were used to classify the HRM algorithm types in our data extraction.

These types of AI algorithms mentioned in the previous paragraph are used to make sense of or find patterns in small or big data sets, as in internal company data (e.g., resumes, employee portfolios, job descriptions, workloads, turnovers, or key performance indicators) and/or massive and diverse datasets from different external data sources (e.g., social media or job search websites) [75, 90]. Precisely, descriptive AI systems are used to analyze, explain, and understand what happened in the past and how it affects the present, such as those used to rank resumes or assess candidate characteristics in the recruitment process [90, 101]. Then, based on past observations, predictive algorithmic systems are used to determine the probability that a situation or behavior will occur in the future, such as those used to predict the future performance of job candidates [90, 101]. Finally, prescriptive systems consider relevant factors and select actions or decisions to be put in place, such as those used to automate candidate screening or suggest candidates to invite to interviews in the recruitment processes [90, 101]. Beyond predicting future outcomes, prescriptive algorithms suggest, decide, or implement actions and decisions in order to support or automate decisions or processes [101]. Overall, extracting both the HRM functions and the algorithm type supported our examination of how distinct

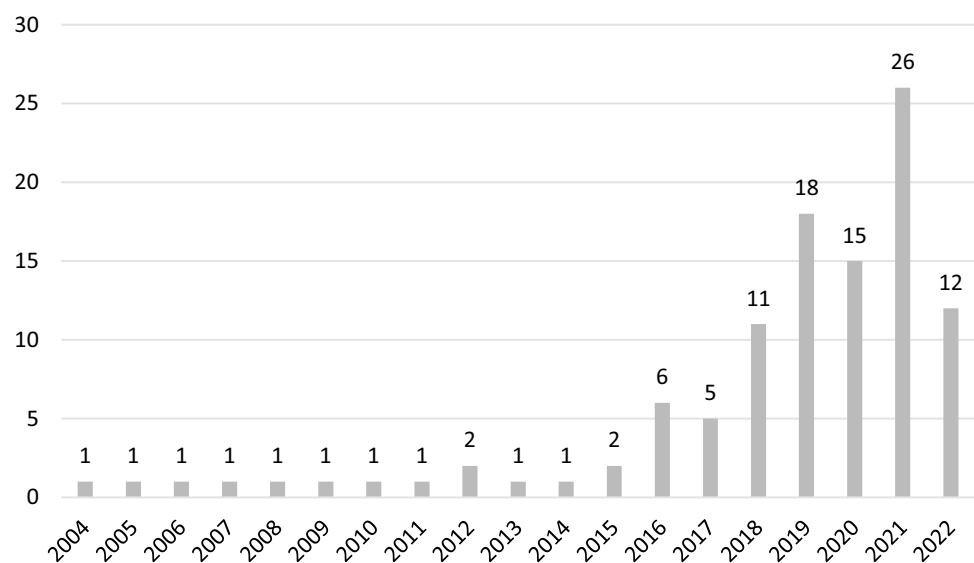
types of AI algorithms are used in each HRM function. This will be elaborated in the next section and detailed examples of how each type of algorithm is used in each HRM function will be provided in the results section.

Regarding the meta-category Responsible AI, we focused on the notion of responsibility related to the use of the system and not those related to the goal of the system. For instance, regardless of whether the finality of the system may be deemed responsible, such as promoting employee wellbeing or sustainable behaviors, our focus was on responsible use of the systems. This is based on the argument that a system with a well-intended goal could still be irresponsible in its use (e.g., a wellbeing system could discriminate against a certain population).

As for the categorisation of responsible (or ethical) AI principles, it is interesting to note that there are over 80 normative frameworks on responsible AI to date around the world [5]. These frameworks present several overlaps and commonalities in principles (e.g., transparency of AI), but also important discrepancies in the terminology (e.g., transparency; explainability; black box; opacity). This fuzziness led us to categorize the responsible AI that emerges from our 107 selected empirical studies according to the most common principles in the responsible AI literature (i.e., autonomy and agentivity, bias and discrimination, explainability and transparency, human role, perceived justice and trust, privacy, system accountability, and working conditions) (primarily based on [5, 6, 13, 16, 19, 58, 100, 128, 147, 152]). We first analyzed whether the studies include responsible practices (Category: Inclusion of responsible practices) and then detailed the practice (Category: Type of responsible practices).

Second, once this step was completed, three members of the research team individually analyzed the summary table, with the goal of identifying points of commonality within categories. The selected empirical peer-reviewed studies varied substantially in terms of vocabulary used, theoretical approaches, aims, disciplines, angles of analysis, and methodologies. This highly diversified sample of studies complexified the analysis of the findings. This led us to adopt an inductive approach for the analysis [4, 122]. This approach aims to generate knowledge about concepts in the literature, rather than validating a pre-existing theory, and the end result comes from generalizing all observations [38, 59].

Guided by the study objectives, we paid particular attention to emerging themes of how AI is currently used in HRM functions (Meta-category: Human Resources Management) and how responsible AI concepts are applied in these empirical studies (Meta-category: Responsible AI). The three researchers then met to compare their analyses. Again, the similarities were strong, and the few dissimilarities were discussed within the whole research team and agreed upon. Notably, regarding the category Human Resource Function,

Fig. 2 Number of studies per year (N = 107)

we enhanced Meijerink and Bondarouk's [101] conceptualization by adding a Health and well-being function because we found several studies falling under this topic. Moreover, regarding the meta-category Responsible AI, only 5 principles emerged from the studies. Those will be elaborated in the next section.

3 Results

3.1 Descriptive results

Our 107 selected empirical and peer-reviewed studies were all published between 2004 and 2022. The median year of publication was 2019. Figure 2 shows the distribution of our 107 empirical studies according to the year they were published. It is important to keep in mind that the year 2022 comprises only the period from January to June, because June 2022 was the month of the data extraction.

In addition, our sample contains 86 different journals or conference proceedings in various fields (e.g., Engineering, Ethics, HRM, Information systems, Management, Mathematics, and Psychology). Table 2 shows the journals or conference proceedings with three or more studies in our sample.

In terms of study design, the selected empirical articles include 63 experimental studies, 15 field studies, 24 studies combining both methods, three case studies and two ethnographies. Moreover, 89 studies used quantitative data, 13 used qualitative data, and five used both. In addition, 69 studies examined the development of a new AI system or model. In the vast majority of these studies, the affordances and the design of the new systems as compared to the old ones were not discussed. Rather, they focused on how their new system

Table 2 Journals or conference proceedings with three or more studies

Name of journal or conference	Count
Computers in Human Behavior	5
Computers & Industrial Engineering	3
Human Resource Management Journal	3
Journal of Applied Psychology	3
Mathematics	3
Proceedings of the ACM Conference on AI, Ethics, and Society	3

offered better validity or performance than past systems or human professionals. In addition, the context surrounding almost all of these developmental studies were laboratory experiments and thus were not implemented and applied to practice.

Regarding the context of all of the 107 studies, they included data collection from 23 different countries: Australia, Bangladesh, Belgium, Canada, China, Colombia, France, Germany, India, Indonesia, Iran, Jordan, Korea, New Zealand, Nigeria, Norway, Palestine, Portugal, Russia, Switzerland, Taiwan, Turkey, and the USA. That said, 29 of the studies in our sample do not specify the country of data collection. The country that recurs the most is the USA (12), and no study includes a cross-country analysis. With respect to sector of activity, 49 studies did not specify a sector under study or it was not applicable. The most studied sector was government or public services (e.g., teachers) with 13 studies, followed by the information technology (IT) sector with 10 studies. The other sectors in our sample are services (7 studies), manufacturing (4), academia (4), power supply (3), military (2),

telecommunication (2), construction (1), sales (1), retail (1) and non-profit organization (1). In addition, nine studies were not specific to the sector studied or reported on a population of workers from various occupations and therefore could not be classified specifically according to their sector of study.

Moreover, many organisations under study were large or multinational organizations (e.g., [10, 26, 98, 119, 149]). This is coherent with the sample sizes of the 69 studies that developed a new AI system or model who often needed to use massive datasets. For example, Avrahami et al. [15] used a longitudinal archival data set comprising of more than 700,000 employees in a large public organization to develop a tool that predicts turnover rates.

3.2 How AI is used in HRM (a technical aspect)

The goal of this section is to report an overview of the affordances of AI in the HRM field based on the empirical studies included in our review. Affordances refer to a use or purpose that a thing can have, that people notice as part of the way they see or experience it. Our 107 selected empirical and peer-reviewed studies include 79 studies that describe how AI is used in specific HRM functions and the types of AI algorithms involved (30 descriptive algorithms, 31 predictive algorithms, and 27 prescriptive algorithms), and 28 that were not specific enough for us to categorize and therefore are not elaborated in this section (e.g., studies on general perception of AI or general use of AI in HRM) (e.g., [7, 21, 43, 64, 69, 81, 82, 120, 135, 150, 151]). Notably, some studies included contain more than one AI algorithm type and/or more than one HRM function.

Table 3 shows the breakdown of HRM algorithm types according to the HRM function. Here, we have augmented the categorization schema proposed by Meijerink and Bondarouk [101] by adding empirically supported uses of AI in each category and by adding the Health and well-being category. Moreover, Fig. 3 shows the distribution of HRM function categories identified in the studies.

3.3 Responsible AI in HRM

Responsible use of AI in HRM encompassed several principles according to our sample of 107 peer-reviewed empirical studies. Six categories emerged from analysis: (1) no responsible principle applied, (2) bias and discrimination, (3) perceived justice and trust, (4) privacy, (5) explainability and transparency, and (6) human role. Some studies applied more than one principle. Appendix 2 shows the classification of studies that clearly applied or investigated responsible AI

principles. Figure 4 shows the distribution of studies across the categories.

3.3.1 No responsible principles reported

Of our sample of 107 empirical studies, 63 did not clearly apply a responsible AI principle. Within these 63 studies, 27 assumed that the AI system would reduce bias and discrimination because it would decrease or eliminate human subjectivity. While this assumption is consistent with some conceptual developments (e.g., [93]), it was not empirically tested in the 27 identified studies.

3.3.2 AI fairness in HRM

The concept of AI fairness in HRM does not seem to have a universally accepted definition across the empirical literature. Instead, we found that it is more of an umbrella term that covers three of our identified principles, namely bias and discrimination in HR-focused AI, perceived justice and trust of decisions and outcomes, and privacy concerns (or intrusiveness) related to AI use.

Twenty studies focused on *detecting or mitigating bias and discrimination* in an AI system for HRM. Indeed, AI-driven decisions can actually be biased and discriminatory because they reflect the data on which they are based [25, 53, 113, 143]. Some studies have looked at how HRM AI tools can be audited and how this auditability may contribute to the detection and mitigation of bias [22, 32, 36, 76, 131, 132, 140, 158]. For example, regarding AI in talent acquisition, Köchling et al. [76] show that AI reproduces (and may even amplify) existing inequalities in the dataset and that underrepresentation of certain groups leads to an unpredictable probability of inviting candidates from those groups to job interviews. Others include the principle of bias by adding a validation step or test in the development of their AI system to demonstrate that the system developed in the study does not discriminate [26, 119, 124, 127]. Finally, the principle of bias and discrimination is applied in empirical research by studies that have developed AI systems whose sole purpose is to detect and mitigate bias and discrimination [12, 60, 124]. For example, Hangartner et al. [60] developed an AI powered tool to continuously monitor hiring discrimination on online recruitment platforms.

Eleven studies in our sample empirically examined the *perception of justice or trust* of decisions and outputs among employees or job seekers, a principle that is also associated with acceptance (e.g., [84]). Most of these studies used an experimental research design in a talent acquisition context and none of them involved the development of a new AI system or model (e.g., [3, 17, 77, 83–85, 89, 111, 141]).

Regarding privacy, only four studies investigated privacy concerns related to AI in HRM [27, 46, 83–85]. Eckhaus

Table 3 Identified HRM functions and types of systems used

HRM algorithm types			
HRM functions	Descriptive	Prescriptive	
Talent acquisition (i.e., actions aimed at recruiting the best possible candidate for a job description)	Development of competency criteria or profile of the best performing candidates [10, 138] Assessment of personality traits, skills, psychological capital or cultural fit of job candidates [8–10, 50, 51, 132] Classification or ranking of resumes or candidates [8, 10, 26, 32, 33, 51, 110, 138, 139] Chatbot assisting job candidates in the process [8, 97, 98] Monitoring hiring discrimination through online recruitment platforms [60]	Predict the potential, suitability, skills, performance and future turnover of job candidates [10, 14, 28, 29, 96, 124, 132, 134, 148, 149] Predict acceptance, conditional acceptance or rejection of candidates [29]	Automated candidate screening or suggestions of which job candidate to invite for job interview or to hire [22, 30, 62, 63, 76, 91, 98, 142, 148, 158] Automated suggestion of matching between job seekers and companies [88]
Performance evaluation (i.e., groups or individuals based on organizational, individual, or pre-established goals)	Detect subjectivity in staff appraisal and reduce bias [1, 42] Generate novel performance evaluation insights [119] Conduct automatic and comprehensive assessments of employees' abilities and competencies [94, 129, 154]	Predict an individual's future performance based on their past performance [36, 78, 80, 94, 95, 98, 104, 110]	Make performance evaluation decisions [17, 47, 48, 68, 89, 92, 98, 111]
Talent management (i.e., management of employee talent present within an organization)	Shed light on career growth patterns [112] Classify an employee's psychological capital [9] Provide an assessment of talent [116, 129]	Predict staffing needs in order to facilitate staff transfer [2] Shed light on the cause and effects in talent management [65] Predict an employee's performance in a given function based on their past or current performance [99, 110]	Resolve issues relating to staff transfer [2] Assign candidates to projects or positions [17, 66, 71, 105] Classify individuals to facilitate talent management [28, 47, 48]
Workforce planning (i.e., job analysis and the prediction of HR demands)	Analyze the turnover phenomenon [15, 106] Assess required skill sets [39, 41, 127]	Predict employee turnover [52, 70, 107, 126, 137, 140, 160, 162] Predict employees at risk of absence leave [86] Predict HR demands to optimize scheduling and HR structures [54]	Enhance the development of strategic HR planning according to turnover predictions [70]
Health and wellbeing (i.e., related to overall mental and physical health of employees)	Recognition of happiness cues [46] Measure of an individual's stress level [55]	Measure cognitive workload [78] Identify employees at risk of sickness [86]	Provide the optimal job rotation scheduling to reduce job stress [136] Provides health-saving strategies [115] Identifies possible areas of job quality enhancement [40] Prescribes personalized work-rest schedules [161]
Compensation and benefits (i.e., monetary and non-monetary)	Detecting gender pay inequalities [12]	Prediction of salary increment [104]	Determine the initial hiring salary of an employee and yearly adjustment [47, 48] Suggestions of adjustments narrowing the pay gap [12]

Some articles contain more than one HRM function and/or HRM algorithm type

[46] shows that scanning emails for data to feed an AI raises privacy concerns, Cayrat and Boxall [27] investigated how organizations implement mechanisms to ensure data privacy and comply with legal obligations (particularly the European General Data Protection Regulation (GDPR)), whereas Langer et al. [83–85] show that the degree of automation of job application process was slightly but positively related to applicants' privacy concerns.

3.3.3 Explainability and transparency in HRM

Explainability (or XAI) is an objective concept as it refers to “[...] an active characteristic of a model, denoting any action or procedure taken by a model with the intent of clarifying or detailing its internal functions” [19, p. 84], while the concept of transparency is more subjective as it can be defined as “[...] the level of awareness and understanding of how [a] system is used” [24, p. 2].

Six studies in our sample applied or investigate this principle [12, 47, 48, 111, 116, 124, 158]. Some studies that have developed AI models have deliberately chosen features that are easier to interpret or provided an explicit explanation of how decisions or outcomes are obtained in order to increase explainability (e.g., [12, 47, 48, 116, 124]). Moreover, Newman et al. [111] directly assessed the effect of this on perceptions of justice in an experimental study by manipulating the level of detail provided about the system process. They found no significant effect.

3.3.4 Human-centered HRM-AI

Finally, 17 studies either applied or investigated the importance of the involvement of humans (e.g., developers, managers, HR practitioners or employees) in the development, implementation and usage of an AI system in HRM. The nature of the human role under study primarily included the level of stakeholders' control over the system (e.g., change or make the final decision, ask questions, appeal, or provide input to the algorithm). The level of users' control or involvement over AI systems seems essential to promote responsible use and even acceptance (e.g., [10, 12, 51, 91, 96, 103, 124]), as “[...] humans must ultimately retain the role of decision makers” [10, p. 66]. For example, Anoaica et al. [12] have put mechanisms in place (mainly in terms of explainability) to provide the HR department with the freedom to make its own judgments and Faliagka et al. [51] warn against blind confidence in an automated system. In a human–computer interaction perspective, as for other domains, providing some degree of control seemed beneficial (e.g., [57]). These findings echo the principle of accountability according to which humans should remain responsible and accountable for their decisions even if supported by AI systems.

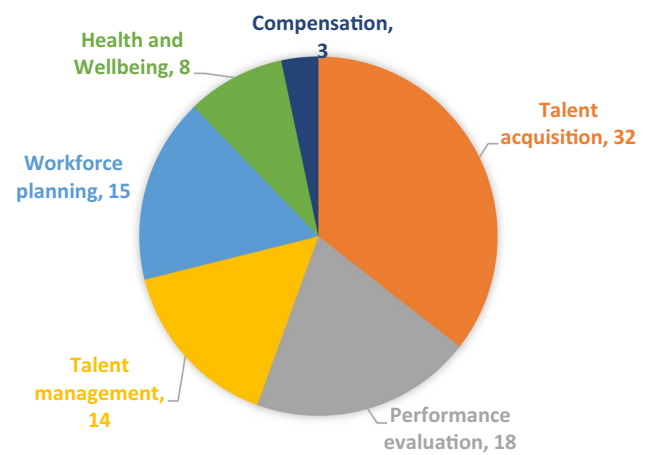


Fig. 3 Distribution of HRM functions in papers with specific AI algorithm type (n=79)

In addition, some studies showed the importance of implication of multiple stakeholders in the development, implementation, and use of AI. In particular, they emphasized that the team be multidisciplinary, continually seek input from a diverse set of stakeholders, and adapt the AI system along the way [97, 148, 149]. The role of HRM in supporting AI systems was also documented. It was mainly addressed through the importance of developing the skills of various stakeholders (e.g., HR practitioners, developers, managers and employees), as multidisciplinary skills are required for the success of AI in HRM [10, 27, 97, 108, 149]. For example, articles highlighted that stakeholders coined as intended users of the systems need to be skilled in statistics and legislation (e.g., GDPR) and understand the responsibility principles surrounding AI, while developers need to be able to go beyond the data and become familiar with HRM [148, 149].

4 Discussion

This paper presents a literature review on empirical and peer-reviewed research on responsible AI in HRM across domains, taking into account the complexity of this phenomenon by looking at both a technical aspect (i.e., how AI is used in HRM) and a social aspect (i.e., responsible AI principles). We contribute to the literature by showing how AI is used in HRM, examine how responsible principles are applied in empirical research of AI in HRM, and evaluate the extent to which these research practices promote responsible AI.

First, our results show that AI in HRM is a multi-domain research topic studied worldwide and across diverse sectors, provided as our sample of 107 empirical and peer-reviewed studies contains 86 different journals or conference

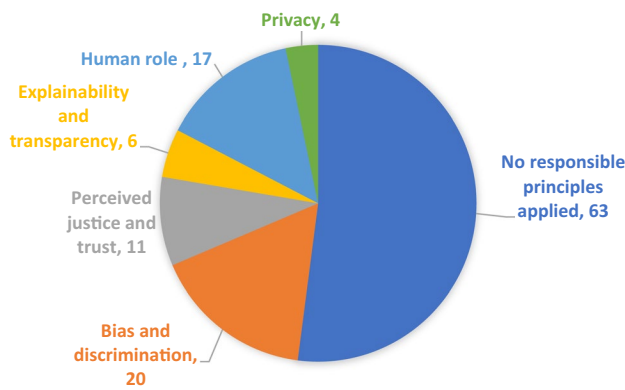


Fig. 4 Distribution of responsible principles across studies (N = 107). Some articles contain more than one principle

proceedings across diverse domains, 23 different countries, and 12 different sectors. Moreover, our descriptive results show that this research topic has greatly increased in popularity over the past decade. Our results also show that the use of three types of AI algorithms (i.e., descriptive, predictive, and prescriptive) have been reported according to six different HRM functions (i.e., 1—talent acquisition, 2—performance evaluation, 3—talent management, 4—workforce planning, 5—health and wellbeing, and 6—compensation) with talent acquisition AI systems being the most empirically studied HRM function and appearing to be the most well implemented. Several explanations can coexist to explain the significant imbalance in terms of interest in the use of AI in different HR functions. For instance, talent acquisition may be more prone to AI systems because it is a task known to be time-consuming, redundant, and subject to human bias, and the data available to train systems includes both actual and potential candidates, so the quantity of data is typically much larger [49].

Our review also highlights that a large number of studies rely on experimental designs or analytical frameworks that have not been tested in real-life settings. Therefore, an important gap with regard to the use of AI in HRM consists of measuring the extent of its effective impact on organizations, the nature of these impacts, as well as the type and size of the organizations concerned. Consequently, the way AI is used in HRM, according to the studies in our sample, and the way it could actually be implemented in organizations may differ. Finally, an important issue that emerges from our analysis concerns the lack of precision regarding the characteristics of the AI tool studied by the researchers, as well as the AI tool's potential context of implementation (organizational and human dimensions), which led to the exclusion of many studies from our review under the selection criteria of clearly being related to human resource management and include an AI-based technology. This gave the impression of a lack

of depth in the literature, which can perhaps be explained by the lack of multi-domain studies on AI in HRM. Indeed, as the studies are mostly carried out in disciplinary silos, they do not allow the researchers to develop a substantially deep and global reflection of the phenomenon. That said, the multi-domain approach of this study provides researchers with perspective, depth, and clarity on which to build, taking into account both the technical and social aspects of AI in HRM.

Second, this paper includes findings on the responsible use of AI in HRM by identifying six categories about how responsible AI is empirically applied and investigated in HRM (i.e., 1—no responsible principle applied, 2—bias and discrimination, 3—perceived justice and trust, 4—privacy, 5—explainability and transparency, and 6—human role). That said, the majority of the studies in our sample did not empirically and clearly examine or incorporate the most common notions of responsibility found in the literature. Therefore, our results show a significant gap between the breadth of conceptual frameworks on responsible AI (e.g., [5, 6, 13, 16, 19, 58, 100, 128, 147, 152]) and the empirical studies that investigate or apply responsible AI principles in HRM. This gap is also observed within the field of HRM, considering the discrepancy between the number of conceptual pieces on the principles surrounding the use of AI in the discipline and the empirical pieces actually examining them. It thus appears that, despite the social and organizational importance of considering the dimensions of responsibility and ethics in the development and use of AI, it is not yet understood and conceptualized as a central dimension in empirical research pertaining to AI systems in HR. We suggest that this state of affairs can be explained by the difficulties involved in integrating the principles of responsibility into empirical research designs. Notably, when empirical research focuses on the effects of tools only once they have been implemented and thus excludes determinants related to the design of the tool itself, it is difficult to identify the full range of possible explanations for gaps in the application of responsibility principles in AI [73, 114].

Considering the elements previously underlined, it is all the more disturbing that some studies have presented AI systems in HRM as more ethical practices than traditional HRM, based on the theoretical argument that AI systems alleviate human subjectivity in practices and therefore reduce bias. This argument is often based on the notion that systems can achieve “fairness through unawareness”, which refers to the ability of systems to not explicitly use protected attributes or to omit sensitive features in the prediction process [31, 35, 79]. However, we found no empirical studies testing whether AI systems are indeed less biased than traditional HRM practices. We would thus discourage any claims that AI systems are less biased than practitioners unless further studies empirically investigate and demonstrate the

validity of this statement. Indeed, we consider that supporting such premises without scientific testing would be irresponsible given that they could wrongly encourage practitioners to adopt AI technologies in order to reduce bias and discrimination.

4.1 Call for future research

Our review clearly shows that both the use of AI in HRM and the application of principles of responsibility are in need of further investigation. Our first and foremost encouragement for future research would be the development of more diversified research protocols that rely on extensive fieldwork and real-life settings. Indeed, as most of the studies in our sample are based on experimental designs, it appears difficult to generalize their contributions to the reality of organizational contexts. Therefore, their contributions for practitioners remain somewhat limited. Moreover, as our results show a major gap between conceptual and empirical research about responsible AI in HRM, we strongly call future research to either apply responsibility principles as a conceptual framework when conducting their empirical work or investigate the effects of responsible principles of AI in HRM. We found little empirical research on topics such as explainability and transparency (in fact, there is no research in our sample on subjective transparency) or privacy. Even the most empirically studied principle in our sample (i.e., bias and discrimination) has received little empirical study relative to the public and academic discussion surrounding it (i.e., [5, 37, 113, 153]). Moreover, perceived justice and/or trust have been primarily studied in experiments and hypothetical scenarios and we call for a methodological diversification such as more research in a real-life context. Regarding the role of humans in responsible AI in HRM, we believe that this principle could be among the most complicated to investigate because the degree and nature of the human role could vary greatly from one situation to another and thus call for more research on this principle. More specifically, although the role of HRM practitioners has been documented, the number of studies was small and knowledge about the outcomes of a high or low role of HRM practitioners on responsible AI remains scarce. In the same vein, although largely discussed in theoretical or conceptual pieces, we know very little on the skills that should be developed among HRM professionals to fully enable them to play this role. We also found that some responsible principles present in the literature were absent from our empirical and peer-reviewed sample. Specifically, our studies did not include empirical research on the impact of AI on stakeholder autonomy and agentivity, or system accountability [100]. We thus call future research on AI in HRM to diversify the approach used to further investigate these responsibility principles.

In addition, we call on future researchers to be explicit and provide as much detail as possible about the AI algorithms being studied, with their characteristics and affordances, as this would allow for a better understanding of how different AI types, features, or responsibility principles affect different outcomes. This could be facilitated through multidisciplinary research teams. In relations to this, we also call on future researchers to take into account the multi-domain nature of responsible AI in HRM by composing multidisciplinary research teams and breaking down silos between research areas. That is, with researchers with advanced technical knowledge of AI as well as researchers with advanced knowledge of HRM. These combinations would allow a better understanding of the complex phenomena of responsible AI in HRM.

As our results show conceptual confusion about responsible principles, we also call for future research to use the knowledge from the conceptual literature and explicitly detail the responsible principle being studied. We found that some empirical studies use terms from AI responsibility such as transparency or discrimination without defining the term or using it in a way that is consistent with the literature. For example, in some studies on AI transparency, this led to conceptual confusion, as researchers were actually studying the concept of explainability.

Also, we found that responsible AI in HRM is studied in many different countries, but we did not find any cross-country analysis. We call for future research to conduct such analyses to further our understanding of responsible AI in HRM and its differences across countries.

Finally, our results show that the field of AI in HRM is evolving rapidly, with the number of studies increasing significantly over the past decade. More empirical work on responsible AI in HRM has already been published since our June 2022 data extraction (e.g., [72]) and we call for future research to continue to update existing reviews.

4.2 Practical implications

For practitioners, our review calls for vigilance in the use of AI within the HRM domain. We have highlighted the lack of research and knowledge about its effects on the workforce. Decision-makers, managers, and HR professionals should be aware of this situation and keep in mind that the benefits of AI for firms also come with risks. Moreover, in order for AI to produce its benefits, it must be carefully crafted and contextualized. Therefore, AI is not a panacea and over-reliance on this technology could come at great costs. This is especially true in the current social context where a high emphasis is placed on issues of equity, diversity, and inclusion (EDI). Among the principles to be considered, the most discussed ones so far are

fairness, explainability and transparency and human role. We also convey policymakers to stay tuned of the future research developments concerning those principles in the elaboration of robust frameworks to regulate to use of AI in HRM.

Appendices

Appendix 1: Search query

("HR" OR "human resource" OR "HRM" OR "human resource management" OR "Human resource management functions" OR "HRM functions" OR "human resource analytics" OR "people analytic" OR "talent analytics" OR "workforce analytics" OR "HR analytics" OR "human capital analytics" OR "Technology-driven HRM" OR "Personnel management" OR "Human Resource management Practices" OR "HRM practices" OR "Talent management" OR "human resources departments" OR "workforce management" OR "HRM decision-making" OR "HRM systems" OR "HR process" OR "HRM role" OR "E-HRM" OR "Human capital" OR "human resources planning" OR "talent management" OR "virtual hrm" OR "Human resource information systems" OR "electronic hrm" OR "HRM systems")

AND

("Responsible" OR "responsible labour practices" OR "responsibility" OR "responsibilities" OR "social responsibility" OR "socially responsible HRM system" OR "ethics" OR "business ethics" OR "ethical work environment" OR "ethical dimension" OR "ethical work climate" OR "ethical decision making" OR "ethical characteristics" OR "organizational ethics" OR "Ethics of labor" OR "Professional ethics" OR "Unethical practices" OR "HRM ethics" OR "ethical climate" OR "ethical dilemmas" OR "ethical organization" OR "Ethical standards" OR "ethical analysis" OR "Values" OR "fairness" OR "discrimination" OR "employment discrimination" OR "Discrimination in employment" OR "diversity" OR "diversity management" OR "inclusion" OR "Decent work" OR "Equality" OR "equality in the workplace" OR "accountability" OR "social inclusion" OR "social integration" OR "organizational inclusion" OR "work conditions" OR "decent work" OR "well-being")

AND

("Artificial intelligence" OR "AI" OR "machine learning" OR "ML" OR "deep learning" OR "RECURRENT neural networks" OR "Artificial Intelligence of Things" OR "DATA mining" OR "SUPERVISED learning" OR "ARTIFICIAL

neural networks" OR "CLASSIFICATION algorithms" OR "Natural language processing" OR "Intelligent automation" OR "Autonomous AI" OR "Chatbots" OR "neural networks" OR "AI tools" OR "Pattern recognition ai" OR "Intelligent Agents" OR "AI applications" OR "Artificial intelligence algorithm").

Appendix 2: Studies that clearly applied responsible AI principles

Principle	Study
Bias and discrimination	[12, 18, 22, 26, 27, 32, 36, 60, 76, 96, 119, 124, 127, 130–132, 140, 148, 157, 158]
Perceived justice and trust	[3, 17, 77, 83–85, 89, 109, 111, 117, 141, 148]
Privacy	[27, 46, 83–85]
Explainability and transparency	[12, 47, 48, 111, 116, 124, 158]
Human role	[10, 12, 27, 51, 57, 62, 63, 91, 96, 97, 103, 106, 108, 124, 148, 149, 158]

Acknowledgements With the financial support of IVADO under the Strategic Research Funding Program: “Human-Centered Artificial Intelligence (HCAI): From Algorithm Development to Human Adoption of AI”, IVADO, PRF-2021-05.

Data availability Not applicable.

Declarations

Conflict of interest On behalf of all authors, the corresponding author states that there is no conflict of interest.

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

1. Abed, A.A., El-Halees, A.M.: Detecting subjectivity in staff performance appraisals by using text mining: teachers appraisals of Palestinian government case study. In: *2017 Palestinian International Conference on Information and Communication Technology (PICICT)*, pp. 120–125. (2017). <https://doi.org/10.1109/picict.2017.25>

2. Acharyya, S., Datta, A.K.: Matching formulation of the Staff Transfer Problem: meta-heuristic approaches. *Opsearch* **57**(3), 629–668 (2020). <https://doi.org/10.1007/s12597-019-00432-w>
3. Acikgoz, Y., Davison, K.H., Compagnone, M., Laske, M.: Justice perceptions of artificial intelligence in selection. *Int. J. Sel. Assess.* **28**(4), 399–416 (2020). <https://doi.org/10.1111/ijssa.12306>
4. Aguinis, H., Ramani, R.S., Alabduljader, N.: Best-practice recommendations for producers, evaluators, and users of methodological literature reviews. *Organ. Res. Methods* **26**(1), 46–76 (2020). <https://doi.org/10.1177/1094428120943281>
5. *AIethicist*. (2022). <https://www.aiethicist.org/ai-principles>
6. Aizenberg, E., van den Hoven, J.: Designing for human rights in AI. *Big Data Soc* (2020). <https://doi.org/10.1177/2053951720949566>
7. Albert, E.T.: AI in talent acquisition: a review of AI-applications used in recruitment and selection. *Strateg. HR Rev.* **18**(5), 215–221 (2019). <https://doi.org/10.1108/shr-04-2019-0024>
8. Allal-Chérif, O., YelaAránega, A., Castaño Sánchez, R.: Intelligent recruitment: how to identify, select, and retain talents from around the world using artificial intelligence. *Technol. Forecast Soc. Change* (2021). <https://doi.org/10.1016/j.techfore.2021.120822>
9. Alola, U.V., Atsa'am, D.D.: Measuring employees' psychological capital using data mining approach. *J. Public Affairs* (2019). <https://doi.org/10.1002/pa.2050>
10. Altemeyer, B.: Making the business case for AI in HR: two case studies. *Strateg. HR Rev.* **18**(2), 66–70 (2019). <https://doi.org/10.1108/shr-12-2018-0101>
11. Angrave, D., Charlwood, A., Kirkpatrick, I., Lawrence, M., Stuart, M.: HR and analytics: why HR is set to fail the big data challenge. *Hum. Resour. Manag. J.* **26**(1), 1–11 (2016). <https://doi.org/10.1111/1748-8583.12090>
12. Anoaica, A., Ben Hassine, A., Deleris, L.A.: Equal pay for equal competences: a statistical approach to address equal pay gap. *ECAI* **2020**, 2949–2955 (2020). <https://doi.org/10.3233/FAIA200468>
13. Ashok, M., Madan, R., Joha, A., Sivarajah, U.: Ethical framework for artificial intelligence and digital technologies. *Int. J. Inf. Manag.* (2022). <https://doi.org/10.1016/j.ijinfomgt.2021.102433>
14. Augusto, D.A., Bernardino, H.S., Barbosa, H.J.C.: Predicting the performance of job applicants by means of genetic programming. In: *2013 BRICS Congress on Computational Intelligence and 11th Brazilian Congress on Computational Intelligence*, 98–103. (2013). <https://doi.org/10.1109/brics-cci-cbic.2013.27>
15. Avrahami, D., Pessach, D., Singer, G., Chalutz Ben-Gal, H.: A human resources analytics and machine-learning examination of turnover: implications for theory and practice. *Int. J. Manpow.* **43**(6), 1405–1424 (2022). <https://doi.org/10.1108/ijm-12-2020-0548>
16. Banks, S.: The ethical use of artificial intelligence in human resource management: a decision-making framework. *Ethics Inf. Technol.* **23**(4), 841–854 (2021). <https://doi.org/10.1007/s10676-021-09619-6>
17. Banks, S., Formosa, P., Griep, Y., Richards, D.: AI decision making with dignity? Contrasting workers' justice perceptions of human and ai decision making in a human resource management context. *Inf. Syst. Front.* **24**(3), 857–875 (2022). <https://doi.org/10.1007/s10796-021-10223-8>
18. Bantilan, N.: Themis-ml: a fairness-aware machine learning interface for end-to-end discrimination discovery and mitigation. *J. Technol. Hum. Serv.* **36**(1), 15–30 (2018). <https://doi.org/10.1080/15228835.2017.1416512>
19. Barredo Arrieta, A., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., García, S., Gil-Lopez, S., Molina, D., Benjamins, R., Chatila, R., Herrera, F.: Explainable Artificial Intelligence (XAI): concepts, taxonomies, opportunities and challenges toward responsible AI. *Inf. Fusion* **58**, 82–115 (2020). <https://doi.org/10.1016/j.inffus.2019.12.012>
20. Basu, S., Majumdar, B., Mukherjee, K., Munjal, S., Palaksha, C.: Artificial intelligence–HRM interactions and outcomes: a systematic review and causal configurational explanation. *Hum. Resour. Manag. Rev.* (2022). <https://doi.org/10.1016/j.hrmr.2022.100893>
21. Bigman, Y.E., Wilson, D., Arnestad, M.N., Waytz, A., Gray, K.: Algorithmic discrimination causes less moral outrage than human discrimination. *J. Exp. Psychol. Gen.* (2022). <https://doi.org/10.1037/xge0001250>
22. Booth, B.M., Hickman, L., Subburaj, S.K., Tay, L., Woo, S.E., D'Mello, S.K.: Integrating psychometrics and computing perspectives on bias and fairness in affective computing: a case study of automated video interviews. *IEEE Signal Process. Mag.* **38**(6), 84–95 (2021). <https://doi.org/10.1109/msp.2021.3106615>
23. Budhwar, P., Malik, A., De Silva, M.T.T., Thevisuthan, P.: Artificial intelligence—challenges and opportunities for international HRM: a review and research agenda. *Int. J. Hum. Resour. Manag.* **33**(6), 1065–1097 (2022). <https://doi.org/10.1080/09585192.2022.2035161>
24. Bujold, A., Parent-Rochelleau, X., Gaudet, M.-C.: Opacity behind the wheel: the relationship between transparency of algorithmic management, justice perception, and intention to quit among truck drivers. *Comput. Hum. Behav. Rep.* **8**, 1–14 (2022). <https://doi.org/10.1016/j.chbr.2022.100245>
25. Buolamwini, J., Gebru, T.: Gender shades: intersectional accuracy disparities in commercial gender classification. In: *Proceedings of the 2018 ACM Conference on Fairness, Accountability and Transparency*, (2018)
26. Campion, M.C., Campion, M.A., Campion, E.D., Reider, M.H.: Initial investigation into computer scoring of candidate essays for personnel selection. *J. Appl. Psychol.* **101**(7), 958–975 (2016). <https://doi.org/10.1037/apl0000108>
27. Cayrat, C., Boxall, P.: Exploring the phenomenon of HR analytics: a study of challenges, risks and impacts in 40 large companies. *J. Organ. Effect. People Perform.* **9**(4), 572–590 (2022). <https://doi.org/10.1108/joepp-08-2021-0238>
28. Chalfin, A., Danieli, O., Hillis, A., Jelveh, Z., Luca, M., Ludwig, J., Mullainathan, S.: Productivity and selection of human capital with machine learning. *Am. Econ. Rev.* **106**(5), 124–127 (2016). <https://doi.org/10.1257/aer.p20161029>
29. Chen, C.-C., Wei, C.-C., Chen, S.-H., Sun, L.-M., Lin, H.-H.: AI predicted competency model to maximize job performance. *Cybern. Syst.* **53**(3), 298–317 (2021). <https://doi.org/10.1080/01969722.2021.1983701>
30. Chen, C.-T., Hung, W.-Z.: A two-phase model for personnel selection based on multi-type fuzzy information. *Mathematics* (2020). <https://doi.org/10.3390/math8101703>
31. Chen, J., Kallus, N., Mao, X., Svacha, G., Udell, M.: Fairness under unawareness. In: *Proceedings of the AMC Conference on Fairness, Accountability, and Transparency*, pp. 339–348. (2019). <https://doi.org/10.1145/3287560.3287594>
32. Chen, L., Ma, R., Hannák, A., Wilson, C.: Investigating the impact of gender on rank in resume search engines. In: *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*, pp. 1–14. (2018). <https://doi.org/10.1145/3173574.3174225>
33. Chien, C.-F., Chen, L.-F.: Data mining to improve personnel selection and enhance human capital: a case study in high-technology industry. *Expert Syst. Appl.* **34**(1), 280–290 (2008). <https://doi.org/10.1016/j.eswa.2006.09.003>
34. Chowdhury, S., Dey, P., Joel-Edgar, S., Bhattacharya, S., Rodriguez-Espindola, O., Abadie, A., Truong, L.: Unlocking the

- value of artificial intelligence in human resource management through AI capability framework. *Hum. Resour. Manag. Rev.* (2022). <https://doi.org/10.1016/j.hrmmr.2022.100899>
35. Cornacchia, G., Anelli, V.W., Biancofiore, G.M., Narducci, F., Pomo, C., Ragone, A., Di Sciascio, E.: Auditing fairness under unawareness through counterfactual reasoning. *Inf. Process. Manag.* (2023). <https://doi.org/10.1016/j.ipm.2022.103224>
 36. Cowgill, B., Dell'Acqua, F., Deng, S., Hsu, D., Verma, N., Chaintreau, A.: Biased programmers? Or biased data? A field experiment in operationalizing AI ethics. In: *Proceedings of the 21st ACM Conference on Economics and Computation.* (2020). <https://doi.org/10.2139/ssrn.3615404>
 37. Crawford, K.: *The Atlas of AI.* Yale University Press (2021)
 38. d'Arripe, A., Oboeuf, A., Routier, C.: L'approche inductive: cinq facteurs propices à son émergence. *Approach. Inductives* **1**(1), 96–124 (2014). <https://doi.org/10.7202/1025747ar>
 39. De Mauro, A., Greco, M., Grimaldi, M., Ritala, P.: Human resources for Big Data professions: a systematic classification of job roles and required skill sets. *Inf. Process. Manag.* **54**(5), 807–817 (2018). <https://doi.org/10.1016/j.ipm.2017.05.004>
 40. del Pozo-Antúnez, J.J., Fernández-Navarro, F., Molina-Sánchez, H., Ariza-Montes, A., Carbonero-Ruz, M.: The machine-part cell formation problem with non-binary values: a MILP model and a case of study in the accounting profession. *Mathematics* (2021). <https://doi.org/10.3390/math9151768>
 41. del Pozo-Antúnez, J.J., Molina-Sánchez, H., Ariza-Montes, A., Fernández-Navarro, F.: Promoting work engagement in the accounting profession: a machine learning approach. *Soc. Indic. Res.* **157**(2), 653–670 (2021). <https://doi.org/10.1007/s11205-021-02665-z>
 42. Dhir, K., Chhabra, A.: Automated employee evaluation using fuzzy and neural network synergism through IoT assistance. *Pers. Ubiquit. Comput.* **23**(1), 43–52 (2019). <https://doi.org/10.1007/s00779-018-1186-6>
 43. Diaz, J.B.B., Young, S.F.: The future is here: a benchmark study of digitally enabled assessment and development tools. *Consult. Psychol. J.* **74**(1), 40–79 (2022). <https://doi.org/10.1037/cpb0000201>
 44. Dick, S.: Artificial intelligence. *Harvard Data Sci. Rev.* (2019). <https://doi.org/10.1162/99608f92.92fe150c>
 45. Duan, Y., Edwards, J.S., Dwivedi, Y.K.: Artificial intelligence for decision making in the era of Big Data—evolution, challenges and research agenda. *Int. J. Inf. Manag.* **48**, 63–71 (2019). <https://doi.org/10.1016/j.ijinfomgt.2019.01.021>
 46. Eckhaus, E.: Measurement of organizational happiness. In: *Advances in Human Factors, Business Management and Leadership*, pp. 266–278. (2018). https://doi.org/10.1007/978-3-319-60372-8_26
 47. Escolar-Jimenez, C.C., Matsuzaki, K., Gustilo, R.C.: A neural-fuzzy network approach to employee performance evaluation. *Int. J. Adv. Trends Comput. Sci. Eng.* **8**(3), 573–581 (2019). <https://doi.org/10.30534/ijatcse/2019/37832019>
 48. Escolar-Jimenez, C.C., Matsuzaki, K., Okada, K., Gustilo, R.C.: Data-driven decisions in employee compensation utilizing a neuro-fuzzy inference system. *Int. J. Emerg. Trends Eng. Res.* **7**(8), 163–169 (2019). <https://doi.org/10.30534/ijeter/2019/10782019>
 49. Eubanks, B.: *Artificial Intelligence for HR: Use AI to Support and Develop a Successful Workforce*, 2nd edn. Kogan Page Publishers, London (2022)
 50. Faliagka, E., Iliadis, L., Karydis, I., Rigou, M., Sioutas, S., Tsakalidis, A., Tzimas, G.: On-line consistent ranking on e-recruitment: seeking the truth behind a well-formed CV. *Artif. Intell. Rev.* **42**(3), 515–528 (2014). <https://doi.org/10.1007/s10462-013-9414-y>
 51. Faliagka, E., Tsakalidis, A., Tzimas, G.: An integrated e-recruitment system for automated personality mining and applicant ranking. *Internet Res.* **22**(5), 551–568 (2012). <https://doi.org/10.1108/10662241211271545>
 52. Fallucchi, F., Coladangelo, M., Giuliano, R., William De Luca, E.: Predicting employee attrition using machine learning techniques. *Computers* (2020). <https://doi.org/10.3390/computers9040086>
 53. Faraj, S., Pachidi, S., Sayegh, K.: Working and organizing in the age of the learning algorithm. *Inf. Organ.* **28**(1), 62–70 (2018). <https://doi.org/10.1016/j.infoandorg.2018.02.005>
 54. Feng, Q., Feng, Z., Su, X.: Design and simulation of human resource allocation model based on double-cycle neural network. *Comput. Intell. Neurosci.* (2021). <https://doi.org/10.1155/2021/7149631>
 55. Freihaut, P., Göritz, A.S.: Using the computer mouse for stress measurement—an empirical investigation and critical review. *Int. J. Hum.-Comput. Stud.* (2021). <https://doi.org/10.1016/j.ijhcs.2020.102520>
 56. Garg, S., Sinha, S., Kar, A.K., Mani, M.: A review of machine learning applications in human resource management. *Int. J. Product. Perform. Manag.* **71**(5), 1590–1610 (2021). <https://doi.org/10.1108/ijppm-08-2020-0427>
 57. Gonzalez, M.F., Liu, W., Shirase, L., Tomczak, D.L., Lobbe, C.E., Justenhoven, R., Martin, N.R.: Allying with AI? Reactions toward human-based, AI/ML-based, and augmented hiring processes. *Comput. Hum. Behav.* (2022). <https://doi.org/10.1016/j.chb.2022.107179>
 58. Goretzko, D., Israel, L.S.F.: Pitfalls of machine learning-based personnel selection. *J. Pers. Psychol.* **21**(1), 37–47 (2022). <https://doi.org/10.1027/1866-5888/a000287>
 59. Guillemette, F.: Approches inductives II. *Recherches Qual.* **28**(2), 1–3 (2009). <https://doi.org/10.7202/1085269ar>
 60. Hangartner, D., Kopp, D., Siegenthaler, M.: Monitoring hiring discrimination through online recruitment platforms. *Nature* **589**(7843), 572–576 (2021). <https://doi.org/10.1038/s41586-020-03136-0>
 61. Herschel, R., Miori, V.M.: Ethics & big data. *Technol. Soc.* **49**, 31–36 (2017). <https://doi.org/10.1016/j.techsoc.2017.03.003>
 62. Hickman, L., Bosch, N., Ng, V., Saef, R., Tay, L., Woo, S.E.: Automated video interview personality assessments: reliability, validity, and generalizability investigations. *J. Appl. Psychol.* **107**(8), 1323–1351 (2022). <https://doi.org/10.1037/apl0000695>
 63. Hickman, L., Saef, R., Ng, V., Woo, S.E., Tay, L., Bosch, N.: Developing and evaluating language-based machine learning algorithms for inferring applicant personality in video interviews. *Hum. Resour. Manag. J.* (2021). <https://doi.org/10.1111/1748-8583.12356>
 64. Hua, Z., Jiang, W., Liang, L.: Adjusting inconsistency through learning in group decision-making, and its application to China's MBA recruiting interview. *Socioecon. Plann. Sci.* **41**(3), 195–207 (2007). <https://doi.org/10.1016/j.seps.2005.08.001>
 65. Huang, L.-C., Huang, K.-S., Huang, H.-P., Jaw, B.-S.: Applying fuzzy neural network in human resource selection system. In: *IEEE Annual Meeting of the Fuzzy Information, 2004. Processing NAFIPS'04*, vol. 1, pp. 169–174. (2004)
 66. Huang, M.-J., Tsou, Y.-L., Lee, S.-C.: Integrating fuzzy data mining and fuzzy artificial neural networks for discovering implicit knowledge. *Knowl.-Based Syst.* **19**(6), 396–403 (2006). <https://doi.org/10.1016/j.knsys.2006.04.003>
 67. Jabotá, M., Santos, J., Gutierrez, I., Moscon, D., Fernandes, P.O., Teixeira, J.P.: Evolution of artificial intelligence research in human resources. *Procedia Comput. Sci.* **164**, 137–142 (2019)
 68. Jing, H.: Application of fuzzy data mining algorithm in performance evaluation of human resource. *Int. Forum Comput.*

- Sci.-Technol. Appl. **2009**, 343–346 (2009). <https://doi.org/10.1109/ifcsta.2009.90>
69. Kaibel, C., Koch-Bayram, I., Biemann, T., Mühlenbock, M.: Applicant perceptions of hiring algorithms-uniqueness and discrimination experiences as moderators. *Acad. Manag. Proc.* (2019). <https://doi.org/10.5465/AMBPP.2019.210>
 70. Kang, I.G., Croft, B., Bichelmeyer, B.A.: Predictors of turnover intention in U.S. Federal Government workforce: machine learning evidence that perceived comprehensive hr practices predict turnover intention. *Public Personnel Manag.* **50**(4), 538–558 (2021). <https://doi.org/10.1177/0091026020977562>
 71. Karatop, B., Kubat, C., Uygün, Ö.: Talent management in manufacturing system using fuzzy logic approach. *Comput. Ind. Eng.* **86**, 127–136 (2015). <https://doi.org/10.1016/j.cie.2014.09.015>
 72. Kassir, S., Baker, L., Dolphin, J., Polli, F.: AI for hiring in context: a perspective on overcoming the unique challenges of employment research to mitigate disparate impact. *AI Ethics* (2022). <https://doi.org/10.1007/s43681-022-00208-x>
 73. Kim, S., Wang, Y., Boon, C.: Sixty years of research on technology and human resource management: looking back and looking forward. *Hum. Resour. Manag.* **60**(1), 229–247 (2020). <https://doi.org/10.1002/hrm.22049>
 74. Kitchin, R.: Big Data, new epistemologies and paradigm shifts. *Big Data Soc.* **1**(1), 1–12 (2014). <https://doi.org/10.1177/2053951714528481>
 75. Kitchin, R., Lauriault, T.P.: Small data in the era of big data. *GeoJournal* **80**(4), 463–475 (2015). <https://doi.org/10.1007/s10708-014-9601-7>
 76. Köchling, A., Riazzy, S., Wehner, M.C., Simbeck, K.: Highly accurate, but still discriminatory. *Bus. Inf. Syst. Eng.* **63**(1), 39–54 (2021). <https://doi.org/10.1007/s12599-020-00673-w>
 77. Köchling, A., Wehner, M.C., Warkocz, J.: Can I show my skills? Affective responses to artificial intelligence in the recruitment process. *Rev. Managerial Sci.* (2022). <https://doi.org/10.1007/s11846-021-00514-4>
 78. Kraft, A.E., Russo, J., Krein, M., Russell, B., Casebeer, W., Ziegler, M.: A systematic approach to developing near real-time performance predictions based on physiological measures. In: *2017 IEEE Conference on Cognitive and Computational Aspects of Situation Management (CogSIMA)*. (2017). <https://doi.org/10.1109/COGSIMA.2017.7929601>
 79. Kusner, M.J., Loftus, J.R.: The long road to fairer algorithms. *Nature* **578**, 34–36 (2020). <https://doi.org/10.1038/d41586-020-00274-3>
 80. Lamarca, B., Ambat, S.: The development of a performance appraisal system using decision tree analysis and fuzzy logic. *Int. J. Intell. Eng. Syst.* **11**(4), 11–19 (2018). <https://doi.org/10.22266/ijies2018.0831.02>
 81. Langer, M., König, C.J., Busch, V.: Changing the means of managerial work: effects of automated decision support systems on personnel selection tasks. *J. Bus. Psychol.* **36**(5), 751–769 (2021). <https://doi.org/10.1007/s10869-020-09711-6>
 82. Langer, M., König, C.J., Fritli, A.: Information as a double-edged sword: the role of computer experience and information on applicant reactions towards novel technologies for personnel selection. *Comput. Hum. Behav.* **81**, 19–30 (2018). <https://doi.org/10.1016/j.chb.2017.11.036>
 83. Langer, M., König, C.J., Hemsing, V.: Is anybody listening? The impact of automatically evaluated job interviews on impression management and applicant reactions. *J. Manag. Psychol.* **35**(4), 271–284 (2020). <https://doi.org/10.1108/jmp-03-2019-0156>
 84. Langer, M., König, C.J., Papathanasiou, M.: Highly automated job interviews: acceptance under the influence of stakes. *Int. J. Sel. Assess.* **27**(3), 217–234 (2019). <https://doi.org/10.1111/ijsa.12246>
 85. Langer, M., König, C.J., Sanchez, D.R.-P., Samadi, S.: Highly automated interviews: applicant reactions and the organizational context. *J. Manag. Psychol.* **35**(4), 301–314 (2020). <https://doi.org/10.1108/jmp-09-2018-0402>
 86. Lawrance, N., Petrides, G., Guerry, M.-A.: Predicting employee absenteeism for cost effective interventions. *Decis. Support Syst.* (2021). <https://doi.org/10.1016/j.dss.2021.113539>
 87. LeCun, Y., Bengio, Y., Hinton, G.: Deep learning. *Nature* **521**(7553), 436–444 (2015). <https://doi.org/10.1038/nature14539>
 88. Lee, D., Ahn, C.: Industrial human resource management optimization based on skills and characteristics. *Comput. Ind. Eng.* **144**, 1–9 (2020). <https://doi.org/10.1016/j.cie.2020.106463>
 89. Lee, M.K.: Understanding perception of algorithmic decisions: fairness, trust, and emotion in response to algorithmic management. *Big Data Soc.* **5**, 1–16 (2018). <https://doi.org/10.1177/2053951718756684>
 90. Leicht-Deobald, U., Busch, T., Schank, C., Weibel, A., Schafheitle, S., Wildhaber, I., Kasper, G.: The challenges of algorithm-based HR decision-making for personal integrity. *J. Bus. Ethics* **160**(2), 377–392 (2019). <https://doi.org/10.1007/s10551-019-04204-w>
 91. Li, L., Lassiter, T., Oh, J., Lee, M.K.: Algorithmic hiring in practice: recruiter and HR professional's perspectives on AI use in hiring. In: *Proceedings of the 2021 ACM Conference on AI, Ethics, and Society*, 166–176. (2021). <https://doi.org/10.1145/3461702.3462531>
 92. Li, Y., Jiang, D., Li, F.: The application of generating fuzzy ID3 algorithm in performance evaluation. *Procedia Eng.* **29**, 229–234 (2012). <https://doi.org/10.1016/j.proeng.2011.12.699>
 93. Lin, Y.-T., Hung, T.-W., Huang, L.T.-L.: Engineering equity: how AI can help reduce the harm of implicit bias. *Philos. Technol.* **34**(S1), 65–90 (2020). <https://doi.org/10.1007/s13347-020-00406-7>
 94. Londoño-Montoya, E., Gomez-Bayona, L., Moreno-López, G., Duarte, C.A., Marín, L.G., Becerra, M.A.: Regression fusion framework: an approach for Human Capital evaluation. In: *Proceedings of the European Conference on Knowledge Management, ECKM, Barcelona, Spain* (2017)
 95. Lopes, S.A., Duarte, M.E., Almeida Lopes, J.: Can artificial neural networks predict lawyers' performance rankings? *Int. J. Product. Perform. Manag.* **67**(9), 1940–1958 (2018). <https://doi.org/10.1108/ijppm-08-2017-0212>
 96. Mahmoud, A.A., Shawabkeh, T.A., Salameh, W.A., Al Amro, I.: Performance predicting in hiring process and performance appraisals using machine learning. In: *2019 10th International Conference on Information and Communication Systems (ICICS)*, pp. 110–115. (2019).
 97. Malik, A., Budhwar, P., Patel, C., Srikanth, N.R.: May the bots be with you! Delivering HR cost-effectiveness and individualised employee experiences in an MNE. *Int. J. Hum. Resour. Manag.* (2020). <https://doi.org/10.1080/09585192.2020.1859582>
 98. Malik, A., De Silva, M.T.T., Budhwar, P., Srikanth, N.R.: Elevating talents' experience through innovative artificial intelligence-mediated knowledge sharing: evidence from an IT-multinational enterprise. *J. Int. Manag.* (2021). <https://doi.org/10.1016/j.intman.2021.100871>
 99. Mallafi, H., Widiantoro, D.H.: Prediction modelling in career management. In: *2016 International Conference on Computational Intelligence and Cybernetics*, pp. 17–21. (2016). <https://doi.org/10.1109/CyberneticsCom.2016.7892560>
 100. Martin, K.: Ethical implications and accountability of algorithms. *J. Bus. Ethics* **160**(4), 835–850 (2018). <https://doi.org/10.1007/s10551-018-3921-3>
 101. Meijerink, J., Bondarouk, T.: The duality of algorithmic management: toward a research agenda on HRM algorithms, autonomy

- and value creation. *Hum. Resour. Manag. Rev.* **33**(1), 1–14 (2023). <https://doi.org/10.1016/j.hrmr.2021.100876>
102. Meijerink, J., Boons, M., Keegan, A., Marler, J.: Algorithmic human resource management: synthesizing developments and cross-disciplinary insights on digital HRM. *Int. J. Hum. Resour. Manag.* **32**(12), 2545–2562 (2021). <https://doi.org/10.1080/09585192.2021.1925326>
 103. Mirowska, A., Mesnet, L.: Preferring the devil you know: potential applicant reactions to artificial intelligence evaluation of interviews. *Hum. Resour. Manag. J.* **32**(2), 364–383 (2021). <https://doi.org/10.1111/1748-8583.12393>
 104. Mobasshera, A., Naher, K., Rezoan Tamal, T.M., Rahman, R.M.: Salary increment model based on fuzzy logic. In: *Artificial Intelligence and Algorithms in Intelligent Systems, Proceedings of 7th Computer Science Online Conference 2018*, vol. 2, pp. 344–353. (2019). https://doi.org/10.1007/978-3-319-91189-2_34
 105. Moon, C., Lee, J., Lim, S.: A performance appraisal and promotion ranking system based on fuzzy logic: an implementation case in military organizations. *Appl. Soft Comput.* **10**(2), 512–519 (2010). <https://doi.org/10.1016/j.asoc.2009.08.035>
 106. Mousavian Anaraki, S.A., Haeri, A., Moslehi, F.: Providing a hybrid clustering method as an auxiliary system in automatic labeling to divide employee into different levels of productivity and their retention. *Iran. J. Manag. Stud.* **15**(2), 207–226 (2022). <https://doi.org/10.22059/IJMS.2021.299705.674004>
 107. Najafi-Zangeneh, S., Shams-Gharneh, N., Arjomandi-Nezhad, A., HashemkhaniZolfani, S.: An improved machine learning-based employees attrition prediction framework with emphasis on feature selection. *Mathematics* (2021). <https://doi.org/10.3390/math9111226>
 108. Nankervis, A., Connell, J., Cameron, R., Montague, A., Priksat, V.: ‘Are we there yet?’ Australian HR professionals and the Fourth Industrial Revolution. *Asia Pac. J. Hum. Resour.* **59**(1), 3–19 (2021). <https://doi.org/10.1111/1744-7941.12245>
 109. Nawaz, N.: Artificial Intelligence interchange human intervention in the recruitment process in Indian Software Industry. *Int. J. Adv. Trends Comput. Sci. Eng.* **8**(4), 1433–1441 (2019). <https://doi.org/10.30534/ijatcse/2019/62842019>
 110. Nedelcu, B.: Human talent forecasting. In *Proceedings of the International Conference on Business Excellence*, vol. 11, no. 1, pp. 437–447. (2017). <https://doi.org/10.1515/picbe-2017-0047>
 111. Newman, D.T., Fast, N.J., Harmon, D.J.: When eliminating bias isn’t fair: algorithmic reductionism and procedural justice in human resource decisions. *Organ. Behav. Hum. Decis. Process.* **160**, 149–167 (2020). <https://doi.org/10.1016/j.obhdp.2020.03.008>
 112. Nikitinsky, N., Kachurina, P., Sergey, S., Shamis, E.: Generation theory in HR practice: text mining for talent management case. In: *Proceedings of the International Conference on Electronic Governance and Open Society: Challenges in Eurasia*, pp. 262–266. (2016). <https://doi.org/10.1145/3014087.3014126>
 113. O’neil, C.: *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Broadway Books, New York (2016)
 114. Orlikowski, W.J., Iacono, C.S.: Desperately seeking the “IT” in IT research—a call to theorizing the IT artifact. *Inf. Syst. Res.* **12**(2), 121–134 (2001)
 115. Orlova, E.: Innovation in company labor productivity management: data science methods application. *Appl. Syst. Innov.* (2021). <https://doi.org/10.3390/asi4030068>
 116. Othman, Z.A., Ismail, N., Nazri, M.Z.A., Jantan, H.: Development of talent model based on publication performance using apriori technique. *Int. J. Adv. Comput. Sci. Appl.* **10**(3), 631–640 (2019). <https://doi.org/10.14569/IJACSA.2019.0100381>
 117. Ötting, S.K., Maier, G.W.: The importance of procedural justice in human-machine interactions: intelligent systems as new decision agents in organizations. *Comput. Hum. Behav.* **89**, 27–39 (2018). <https://doi.org/10.1016/j.chb.2018.07.022>
 118. Page, M.J., McKenzie, J.E., Bossuyt, P.M., Boutron, I., Hoffmann, T.C., Mulrow, C.D., Shamseer, L., Tetzlaff, J.M., Akl, E.A., Brennan, S.E., Chou, R., Glanville, J., Grimshaw, J.M., Hrobjartsson, A., Lalu, M.M., Li, T., Loder, E.W., Mayo-Wilson, E., McDonald, S., McGuinness, L.A., Stewart, L.A., Thomas, J., Tricco, A.C., Welch, V.A., Whiting, P., Moher, D.: The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *BMJ* **372**, n71 (2021). <https://doi.org/10.1136/bmj.n71>
 119. Palshikar, G.K., Apte, M., Pawar, S., Ramrakhiyani, N.: HiSPEED: a system for mining performance appraisal data and text. *Int. J. Data Sci. Anal.* **8**(1), 95–111 (2019). <https://doi.org/10.1007/s41060-018-0142-x>
 120. Pan, Y., Froese, F., Liu, N., Hu, Y., Ye, M.: The adoption of artificial intelligence in employee recruitment: the influence of contextual factors. *Int. J. Hum. Resour. Manag.* (2021). <https://doi.org/10.1080/09585192.2021.1879206>
 121. Pan, Y., Froese, F.J.: An interdisciplinary review of AI and HRM: challenges and future directions. *Hum. Resour. Manag. Rev.* (2022). <https://doi.org/10.1016/j.hrmr.2022.100924>
 122. Paré, G., Trudel, M.-C., Jaana, M., Kitsiou, S.: Synthesizing information systems knowledge: a typology of literature reviews. *Inf. Manag.* **52**(2), 183–199 (2015). <https://doi.org/10.1016/j.im.2014.08.008>
 123. Pereira, V., Hadjielias, E., Christofi, M., Vrontis, D.: A systematic literature review on the impact of artificial intelligence on workplace outcomes: a multi-process perspective. *Hum. Resour. Manag. Rev.* (2021). <https://doi.org/10.1016/j.hrmr.2021.100857>
 124. Pessach, D., Singer, G., Avrahami, D., Chalutz Ben-Gal, H., Shmueli, E., Ben-Gal, I.: Employees recruitment: a prescriptive analytics approach via machine learning and mathematical programming. *Decis. Support Syst.* **134**, 1–18 (2020). <https://doi.org/10.1016/j.dss.2020.113290>
 125. Priksat, V., Malik, A., Budhwar, P.: AI-augmented HRM: antecedents, assimilation and multilevel consequences. *Hum. Resour. Manag. Rev.* (2021). <https://doi.org/10.1016/j.hrmr.2021.100860>
 126. Punnoose, R., Ajit, P.: Prediction of employee turnover in organizations using machine learning algorithms. *Int. J. Adv. Res. Artif. Intell.* **5**(9), 22–26 (2016). <https://doi.org/10.14569/IJARAI.2016.050904>
 127. Putka, D.J., Oswald, F.L., Landers, R.N., Beatty, A.S., McCloy, R.A., Yu, M.C.: Evaluating a natural language processing approach to estimating KSA and interest job analysis ratings. *J. Bus. Psychol.* (2022). <https://doi.org/10.1007/s10869-022-09824-0>
 128. Qamar, Y., Agrawal, R.K., Samad, T.A., Chiappetta Jabbour, C.J.: When technology meets people: the interplay of artificial intelligence and human resource management. *J. Enterp. Inf. Manag.* **34**(5), 1339–1370 (2021). <https://doi.org/10.1108/jeim-11-2020-0436>
 129. Quan, P., Liu, Y., Zhang, T., Wen, Y., Wu, K., He, H., Shi, Y.: A novel data mining approach towards human resource performance appraisal. In: *International Conference on Computational Science—ICCS 2018*, pp. 476–488. (2018). https://doi.org/10.1007/978-3-319-93701-4_37
 130. Raghavan, M., Barocas, S., Kleinberg, J., Levy, K.: Mitigating bias in algorithmic hiring. In: *Proceedings of the 2020 ACM Conference on Fairness, Accountability, and Transparency*, pp. 469–481. (2020). <https://doi.org/10.1145/3351095.3372828>
 131. Ramezanzadehmoghadam, M., Chi, H., Jones, E.L., Chi, Z.: Inherent discriminability of BERT towards racial minority associated data. In: *International Conference on Computational Science and Its Applications*, vol. 12951, pp. 256–271. (2021). https://doi.org/10.1007/978-3-030-86970-0_19

132. Rhea, A., Markey, K., D'Arinzo, L., Schellmann, H., Sloane, M., Squires, P., Stoyanovich, J.: Resume format, linkedin URLs and other unexpected influences on AI personality prediction in hiring: results of an audit. In: *Proceedings of the 2022 ACM Conference on AI, Ethics, and Society*, 572–587. (2022). <https://doi.org/10.1145/3514094.3534189>
133. Rodgers, W., Murray, J.M., Stefanidis, A., Degbey, W.Y., Tarba, S.Y.: An artificial intelligence algorithmic approach to ethical decision-making in human resource management processes. *Hum. Resour. Manag. Rev.* (2023). <https://doi.org/10.1016/j.hrmr.2022.100925>
134. Sajjadi, S., Sojourner, A.J., Kammeyer-Mueller, J.D., Mykerezi, E.: Using machine learning to translate applicant work history into predictors of performance and turnover. *J. Appl. Psychol.* **104**(10), 1207–1225 (2019). <https://doi.org/10.1037/apl0000405>
135. Schick, J., Fischer, S.: Dear computer on my desk, which candidate fits best? An assessment of candidates' perception of assessment quality when using AI in personnel selection. *Front. Psychol.* **12**, 1–11 (2021). <https://doi.org/10.3389/fpsyg.2021.739711>
136. Sebt, V., Ghasemi, S.S.: Presenting a comprehensive smart model of job rotation as a corporate social responsibility to improve human capital. *Int. J. Supply Oper. Manag.* **8**(2), 212–231 (2021). <https://doi.org/10.22034/IJSOM.2021.2.7>
137. Sexton, R.S., McMurtrey, S., Michalopoulos, J.O., Smith, A.M.: Employee turnover: a neural network solution. *Comput. Oper. Res.* **32**(10), 2635–2651 (2005). <https://doi.org/10.1016/j.cor.2004.06.022>
138. Shahhosseini, V., Sebt, M.: Competency-based selection and assignment of human resources to construction projects. *Scientia Iranica* **18**(2), 163–180 (2011). <https://doi.org/10.1016/j.scient.2011.03.026>
139. Shehu, M.A., Saeed, F.: An adaptive personnel selection model for recruitment using domain-driven data mining. *J. Theor. Appl. Inf. Technol.* **91**(1), 117–130 (2016)
140. Speer, A.B.: Empirical attrition modelling and discrimination: balancing validity and group differences. *Hum. Resour. Manag. J.* (2021). <https://doi.org/10.1111/1748-8583.12355>
141. Suen, H.-Y., Chen, M.Y.-C., Lu, S.-H.: Does the use of synchrony and artificial intelligence in video interviews affect interview ratings and applicant attitudes? *Comput. Hum. Behav.* **98**, 93–101 (2019). <https://doi.org/10.1016/j.chb.2019.04.012>
142. Suen, H.-Y., Hung, K.-E., Lin, C.-L.: Intelligent video interview agent used to predict communication skill and perceived personality traits. *Hum.-centric Comput. Inf. Sci.* (2020). <https://doi.org/10.1186/s13673-020-0208-3>
143. Tambe, P., Cappelli, P., Yakubovich, V.: Artificial intelligence in human resources management: challenges and a path forward. *Calif. Manag. Rev.* **61**(4), 15–42 (2019). <https://doi.org/10.1177/0008125619867910>
144. Tegmark, M.: *Life 3.0: Being Human in the Age of Artificial Intelligence*. Vintage, New York (2017)
145. Trenerry, B., Chng, S., Wang, Y., Suhaila, Z.S., Lim, S.S., Lu, H.Y., Oh, P.H.: Preparing workplaces for digital transformation: an integrative review and framework of multi-level factors. *Front. Psychol.* (2021). <https://doi.org/10.3389/fpsyg.2021.620766>
146. Tricco, A.C., Lillie, E., Zarin, W., O'Brien, K., Colquhoun, H., Kastner, M., Levac, D., Ng, C., Sharpe, J.P., Wilson, K., Kenny, M., Warren, R., Wilson, C., Stelfox, H.T., Straus, S.E.: A scoping review on the conduct and reporting of scoping reviews. *BMC Med. Res. Methodol.* **16**, 15 (2016). <https://doi.org/10.1186/s12874-016-0116-4>
147. Tursunbayeva, A., Pagliari, C., Di Lauro, S., Antonelli, G.: The ethics of people analytics: risks, opportunities and recommendations. *Pers. Rev.* **51**(3), 900–921 (2021). <https://doi.org/10.1108/pr-12-2019-0680>
148. van den Broek, E., Sergeeva, A., Huysman, M.: Hiring algorithms: an ethnography of fairness in practice. In: *ICIS 2019 Proceedings*, vol. 6. (2019). https://aisel.aisnet.org/icis2019/future_of_work/future_work/6
149. van den Broek, E., Sergeeva, A., Huysman Vrije, M.: When the machine meets the expert: an ethnography of developing AI for hiring. *MIS Q.* **45**(3), 1557–1580 (2021). <https://doi.org/10.25300/misq/2021/16559>
150. van Esch, P., Black, J.S., Arli, D.: Job candidates' reactions to AI-enabled job application processes. *AI Ethics* **1**(2), 119–130 (2021). <https://doi.org/10.1007/s43681-020-00025-0>
151. van Esch, P., Black, J.S., Ferolie, J.: Marketing AI recruitment: the next phase in job application and selection. *Comput. Hum. Behav.* **90**, 215–222 (2019). <https://doi.org/10.1016/j.chb.2018.09.009>
152. Varma, A., Dawkins, C., Chaudhuri, K.: Artificial intelligence and people management: a critical assessment through the ethical lens. *Hum. Resour. Manag. Rev.* (2022). <https://doi.org/10.1016/j.hrmr.2022.100923>
153. Vassilopoulou, J., Kyriakidou, O., Özbilgin, M.F., Groutsis, D.: Scientism as illusio in HR algorithms: towards a framework for algorithmic hygiene for bias proofing. *Hum. Resour. Manag. J.* (2022). <https://doi.org/10.1111/1748-8583.12430>
154. Wang, J., Lin, Y.-I., Hou, S.-Y.: A data mining approach for training evaluation in simulation-based training. *Comput. Ind. Eng.* **80**, 171–180 (2015). <https://doi.org/10.1016/j.cie.2014.12.008>
155. Wang, P.: On defining artificial intelligence. *J. Artif. Gen. Intell.* **10**(2), 1–37 (2019). <https://doi.org/10.2478/jagi-2019-0002>
156. Webster, J., Watson, R.T.: Analyzing the past to prepare for the future: writing a literature review. *MIS Q.* **26**(2), xiii–xxiii (2002)
157. Williams, S.D.: A textual analysis of racial considerations in human resource analytics vendors' marketing. *Manag. Res. Pract.* **12**(4), 49–63 (2020)
158. Wilson, C., Ghosh, A., Jiang, S., Mislove, A., Baker, L., Szary, J., Trindel, K., Polli, F.: Building and auditing fair algorithms. In: *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, pp. 666–677. (2021). <https://doi.org/10.1145/3442188.3445928>
159. Wohlin, C.: Guidelines for snowballing in systematic literature studies and a replication in software engineering. In: *Proceedings of the 18th International Conference on Evaluation and Assessment in Software Engineering*, pp. 1–10. (2014). <https://doi.org/10.1145/2601248.2601268>
160. Yiğit, İ. O., Shourabizadeh, H.: An approach for predicting employee churn by using data mining. In: *International Artificial Intelligence and Data Processing Symposium*. (2017). <https://ieeexplore-ieee.org.proxy2.hec.ca/stamp/stamp.jsp?tp=&arnumber=8090324>
161. Yu, H., Miao, C., Zheng, Y., Cui, L., Fauvel, S., Leung, C.: Ethically aligned opportunistic scheduling for productive laziness. In: *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, pp. 45–51. (2019). <https://doi.org/10.1145/3306618.3314240>
162. Zhao, Y., Hryniewicki, M. K., Cheng, F., Fu, B., Zhu, X.: Employee turnover prediction with machine learning: a reliable approach. In: *Intelligent Systems and Applications*, pp. 737–758. (2018). https://doi.org/10.1007/978-3-030-01057-7_56

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.